

# Introduction to Reinforcement Learning based Control and Deep Reinforcement Learning based control



## Research

Centre de Recherche en Automatique de Nancy  
(CRAN) UMR 7039,  
Faculté des Sciences et Technologies  
Boulevard des Aiguillettes - BP 70239 - Bât.  
1er cycle 54506 Vandoeuvre-lès-Nancy  
Cedex France

Associate Professor  
(Maitre de Conférences)

Automatic Control, Reliability of Systems  
[mayank-shekhar.jha@univ-lorraine.fr](mailto:mayank-shekhar.jha@univ-lorraine.fr)



## Teaching

Polytech Nancy (ESSTIN),  
2 Rue Jean Lamour 54509  
Vandoeuvre-lès-Nancy,  
Cedex France



Email: [mayank-shekhar.jha@univ-lorraine.fr](mailto:mayank-shekhar.jha@univ-lorraine.fr)

Introduction

Markov Decision process

Dynamic Programming for planning  
and control

# Dynamic Programming

# Dynamic Programming

## Policy Evaluation

Iterative approach required

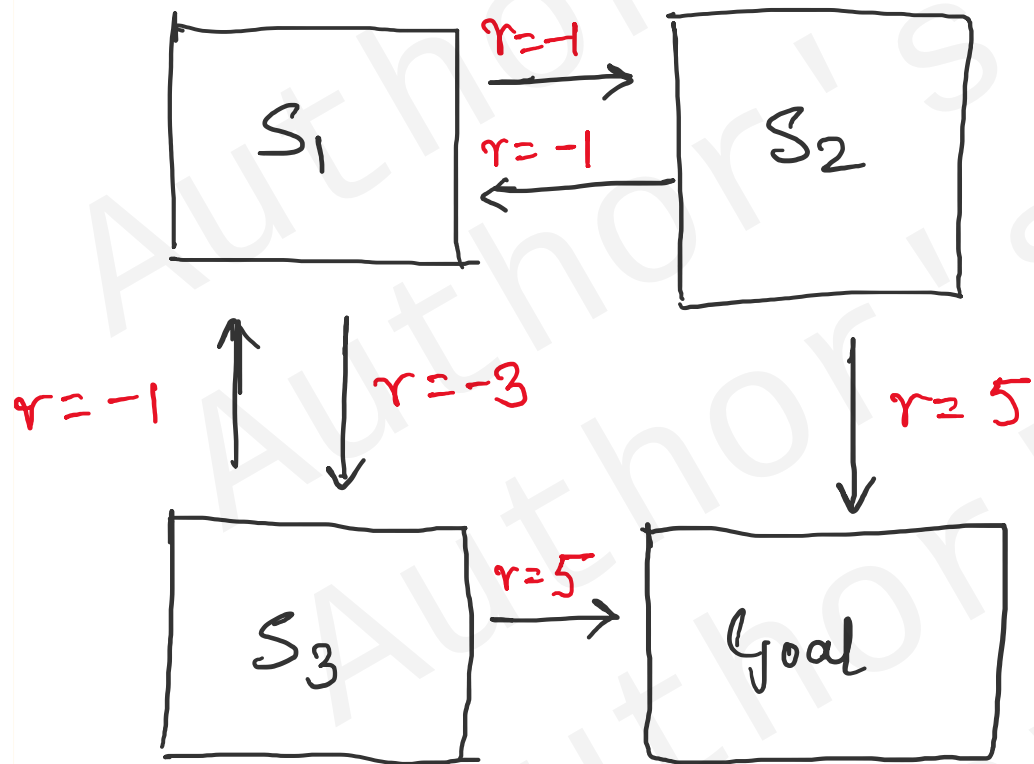
# Dynamic Programming

## Policy Evaluation

Iterative approach required

Motivation.

# Grid World Example



Given  $\pi$ ,  
determine  $V_\pi(s) \forall s \in S$

→ for all

Policy Example:  $\pi$

Actions are drawn from a uniform distribution.

$$\pi(\text{right} | s_1) = \frac{1}{2}$$

$$\pi(\text{down} | s_1) = \frac{1}{2}$$

$$\pi(\text{left} | s_2) = \frac{1}{2}$$

$$\pi(\text{down} | s_2) = \frac{1}{2}$$

$$\pi(\text{up} | s_2) = \frac{1}{2}$$

$$\pi(\text{right} | s_2) = \frac{1}{2}$$

Actions are based on policy.

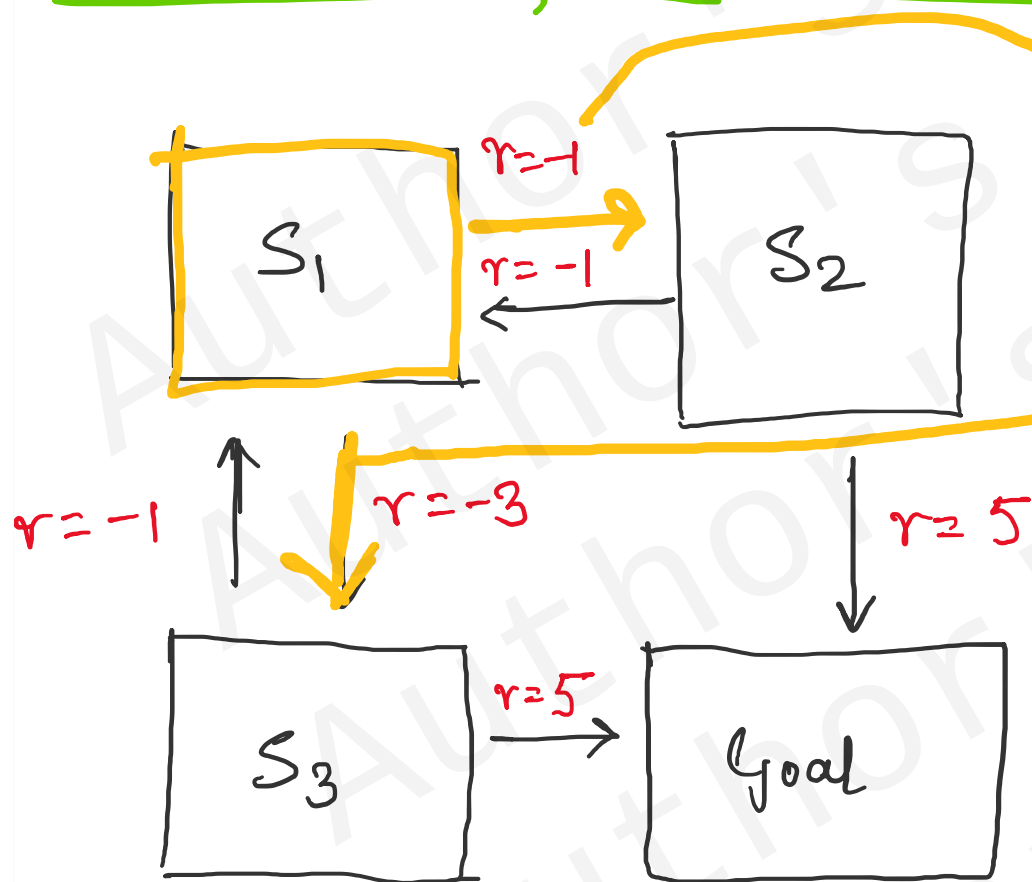
Policy depends on some probability

# Grid World Example

Half time

Other half time

$$V_{\pi}(s_1) = \frac{1}{2} (-1 + V_{\pi}(s_2)) + \frac{1}{2} (-3 + V_{\pi}(s_3))$$

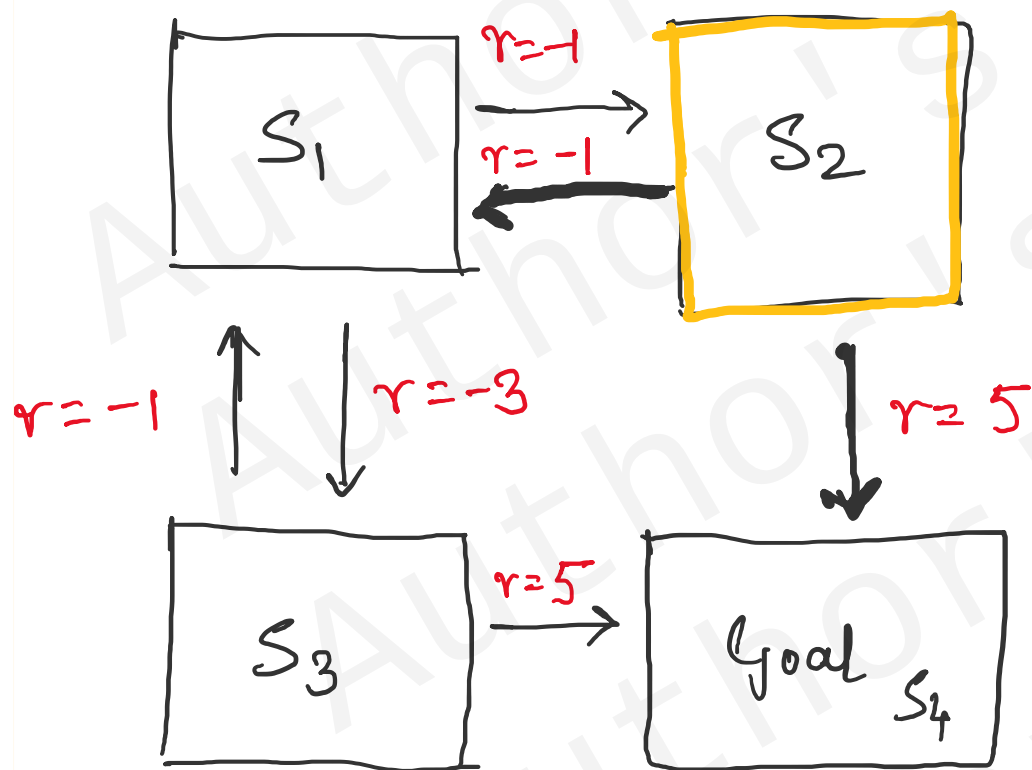


Given  $\pi$ , determine  $V_{\pi}(s) \forall s \in S$  → Random Uniform Policy.

# Grid World Example

Half time

Other half time



$$V_{\pi}(S_1) = \frac{1}{2} (-1 + V_{\pi}(S_2)) + \frac{1}{2} (-3 + V_{\pi}(S_3))$$

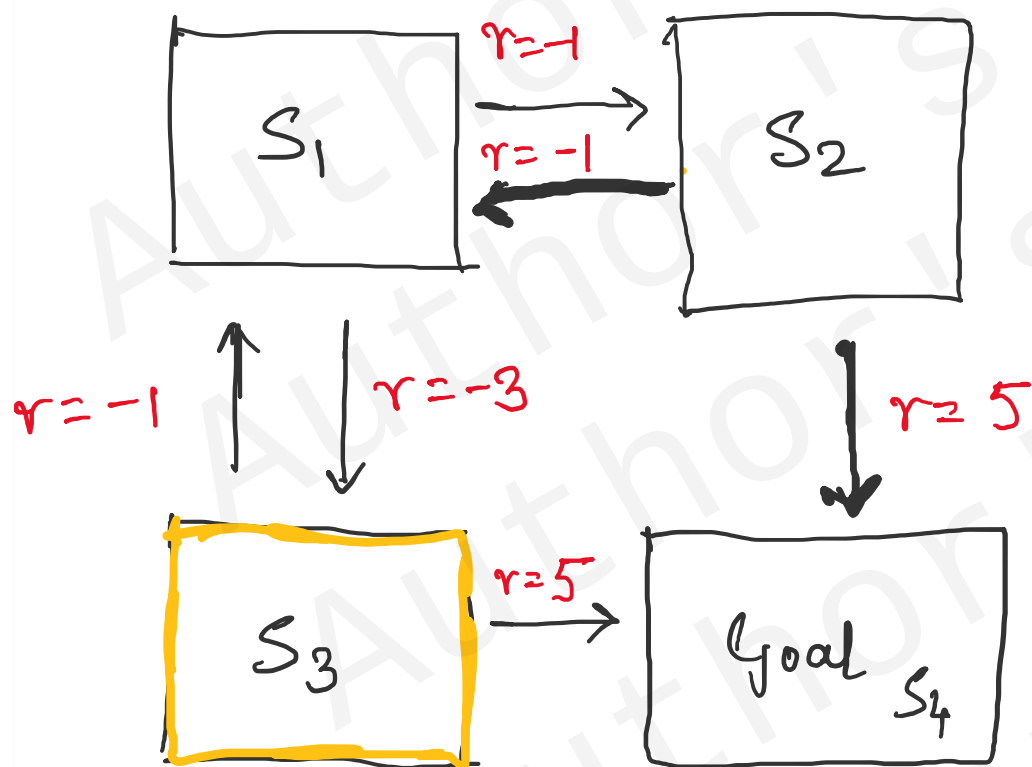
$$V_{\pi}(S_2) = \frac{1}{2} (-1 + V_{\pi}(S_1)) + \frac{1}{2} (5 + V_{\pi}(S_4))$$

Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

# Grid World Example

Half time

Other half time



$$V_{\pi}(s_1) = \frac{1}{2} (-1 + V_{\pi}(s_2)) + \frac{1}{2} (-3 + V_{\pi}(s_3))$$

$$V_{\pi}(s_2) = \frac{1}{2} (-1 + V_{\pi}(s_1)) + \frac{1}{2} (5 + V_{\pi}(s_4))$$

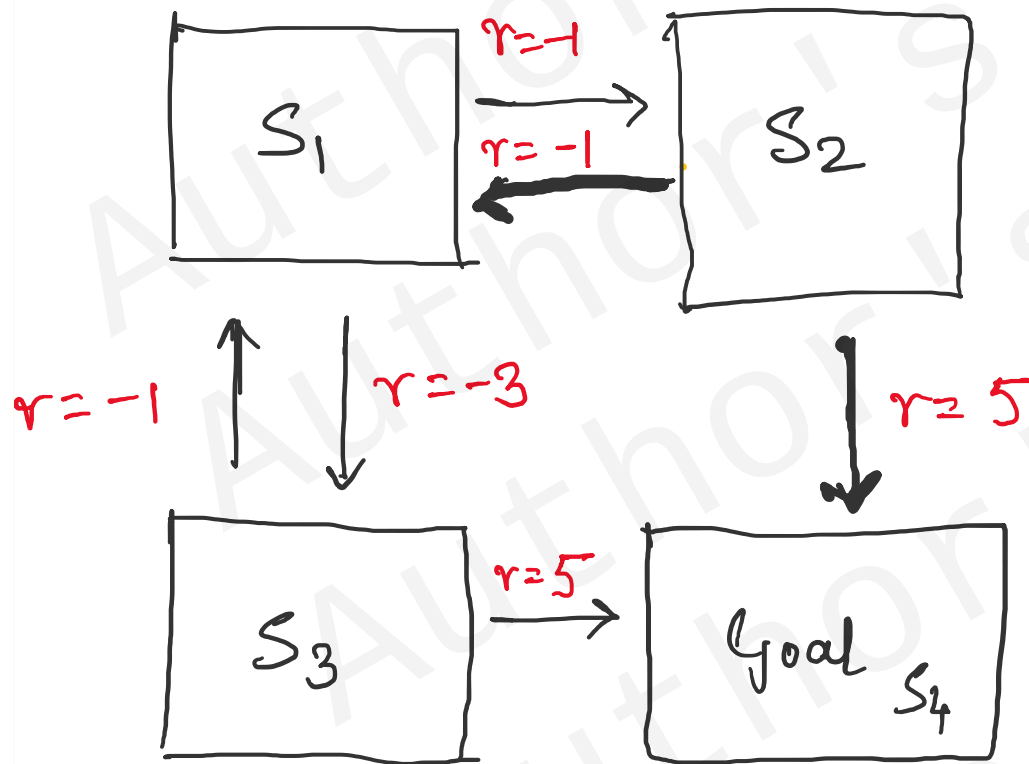
$$V_{\pi}(s_3) = \frac{1}{2} (-1 + V_{\pi}(s_1)) + \frac{1}{2} (5 + V_{\pi}(s_4))$$

Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

# Grid World Example

Half time

Other half time



$$V_{\pi}(s_1) = \frac{1}{2} (-1 + V_{\pi}(s_2)) + \frac{1}{2} (-3 + V_{\pi}(s_3))$$

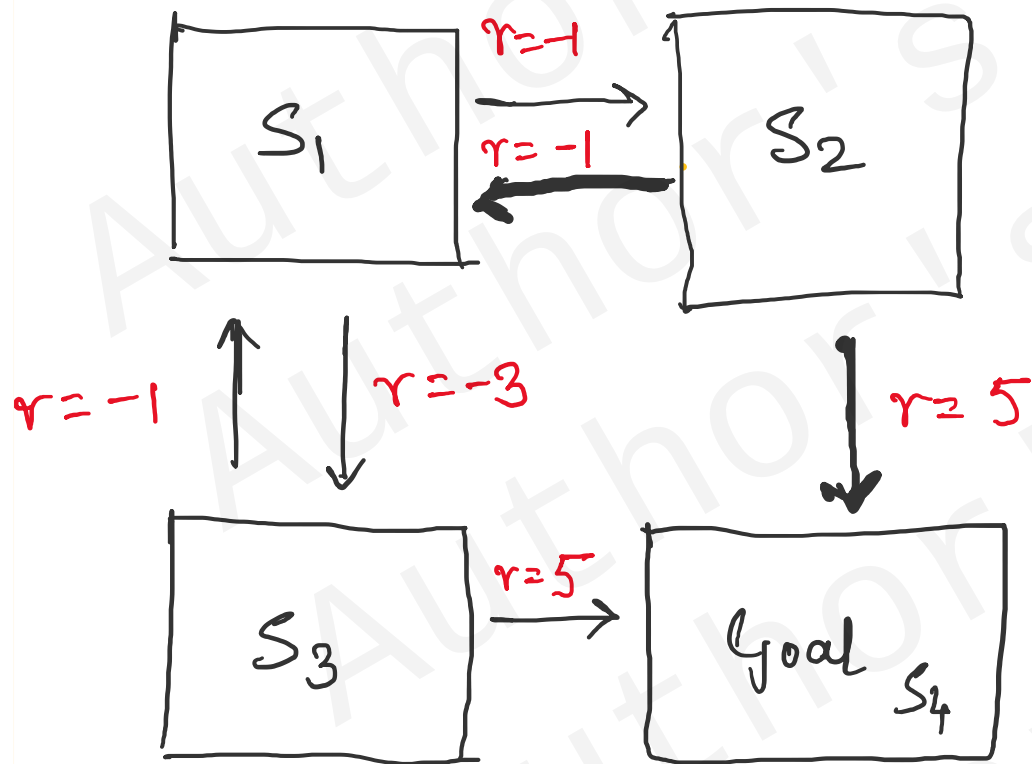
$$V_{\pi}(s_2) = \frac{1}{2} (-1 + V_{\pi}(s_1)) + \frac{1}{2} (5 + V_{\pi}(s_4))$$

$$V_{\pi}(s_3) = \frac{1}{2} (-1 + V_{\pi}(s_1)) + \frac{1}{2} (5 + V_{\pi}(s_4))$$

$$V_{\pi}(s_4) = 0$$

Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

# Grid World Example



Given  $\pi$ ,  
determine  $V_\pi(s) \forall s \in S$

$$V_\pi(s_1) = \frac{1}{2} (-1 + V_\pi(s_2)) + \frac{1}{2} (-3 + V_\pi(s_3))$$

$$V_\pi(s_2) = \frac{1}{2} (-1 + V_\pi(s_1)) + \frac{1}{2} (5 + V_\pi(s_4))$$

$$V_\pi(s_3) = \frac{1}{2} (-1 + V_\pi(s_1)) + \frac{1}{2} (5 + V_\pi(s_4))$$

$$V_\pi(s_4) = 0$$

4 equations, 4 unknown variables

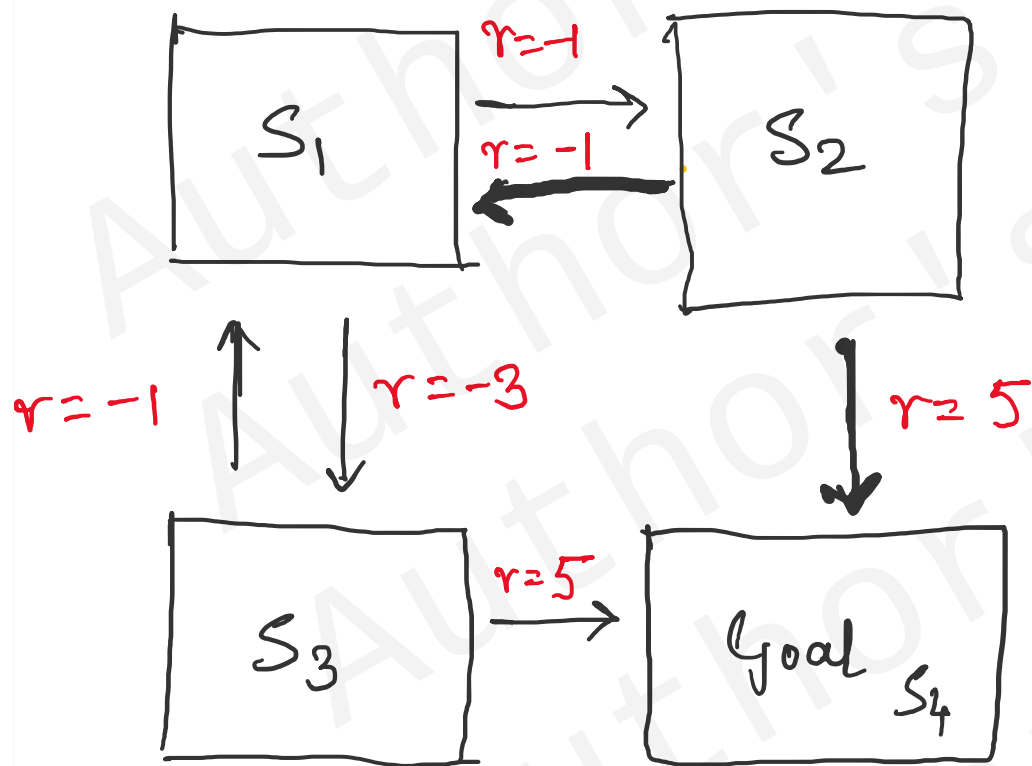
Can be solved to give

$$V_\pi(s_1) = 0, V_\pi(s_4) = 0$$

$$V_\pi(s_2) = 2$$

$$V_\pi(s_3) = 2$$

# Grid World Example



Given  $\pi$ ,  
determine  $V_\pi(s) \forall s \in S$

$$V_\pi(s_1) = \frac{1}{2} (-1 + V_\pi(s_2)) + \frac{1}{2} (-3 + V_\pi(s_3))$$

$$V_\pi(s_2) = \frac{1}{2} (-1 + V_\pi(s_1)) + \frac{1}{2} (5 + V_\pi(s_4))$$

$$V_\pi(s_3) = \frac{1}{2} (-1 + V_\pi(s_1)) + \frac{1}{2} (5 + V_\pi(s_4))$$

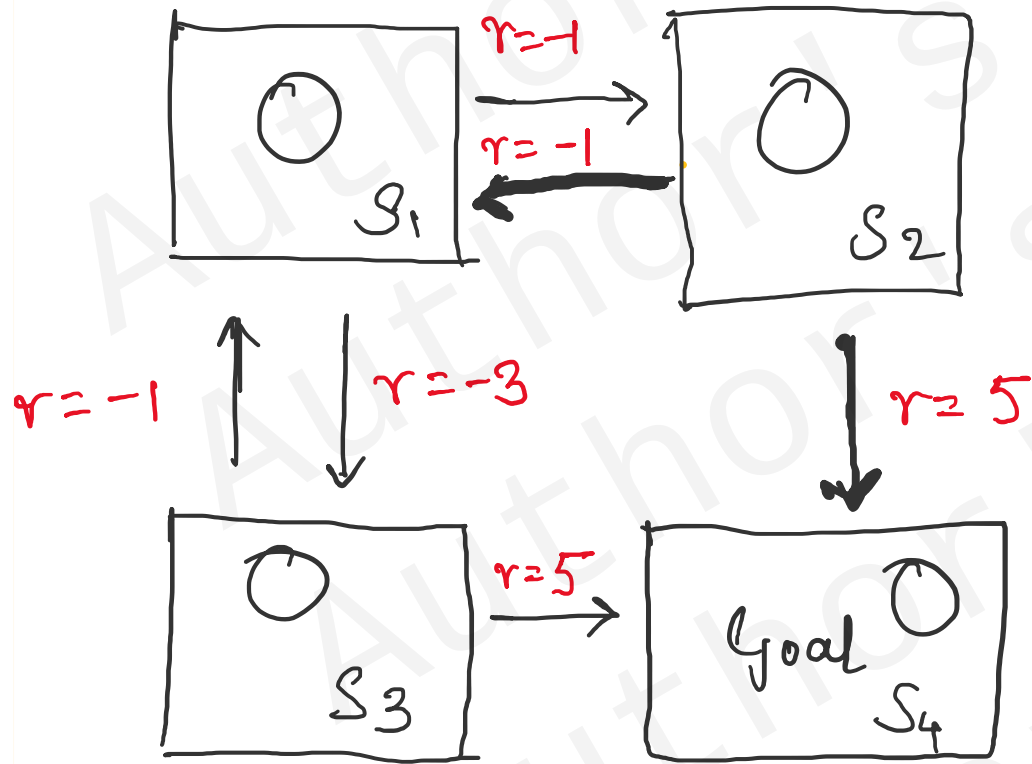
$$V_\pi(s_4) = 0$$

How to handle large state space?

How to optimise under several policies / large states??

Iterative Approach.

# Grid World Example (Iterative Approach)



Given  $\Pi$ ,  
determine  $V_{\Pi}(s) \forall s \in S$

→ start with a guess  
 $\{s_1=0, s_2=0, s_3=0, s_4=0\}$

→ Consider one state  $S_1$

→ Improve the guess with respect to that state

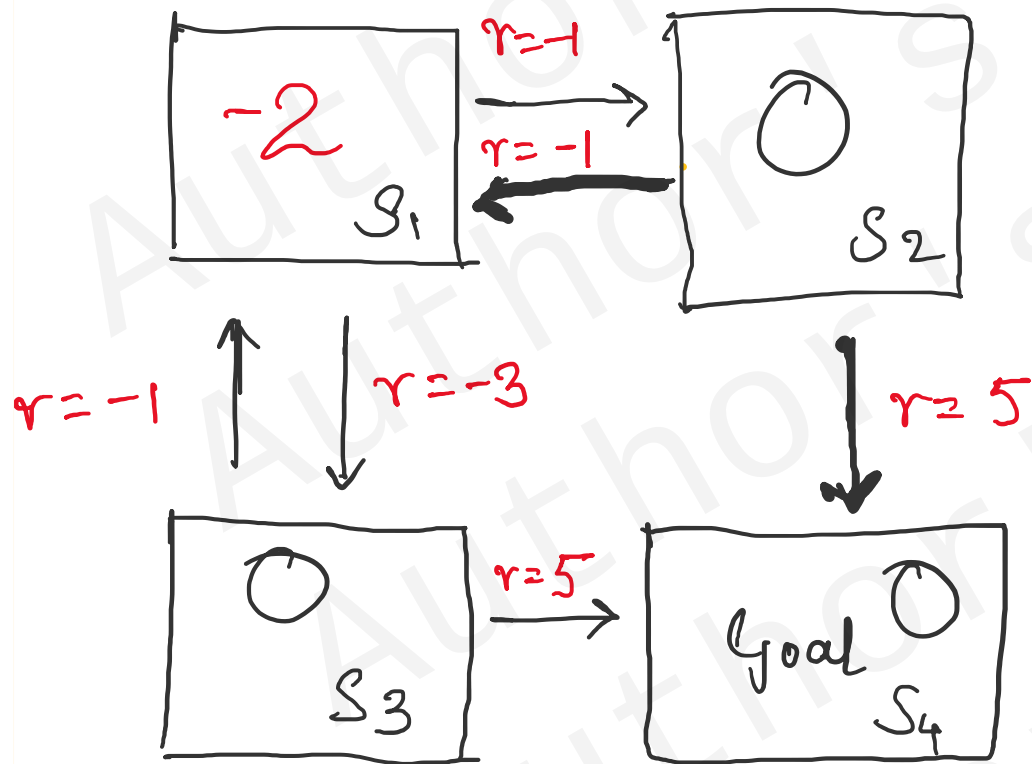
$$V_{\Pi}(S_1) = \frac{1}{2}(-1 + V_{\Pi}(S_2)) + \frac{1}{2}(-3 + V_{\Pi}(S_3))$$

Use this to update the guess  $V_{\Pi}(S_1)$

update guess with a guess

$$V_{\Pi}(S_1) \leftarrow \frac{1}{2}(-1 + 0) + \frac{1}{2}(-3 + 0) = -2$$

# Grid World Example (Iterative Approach)

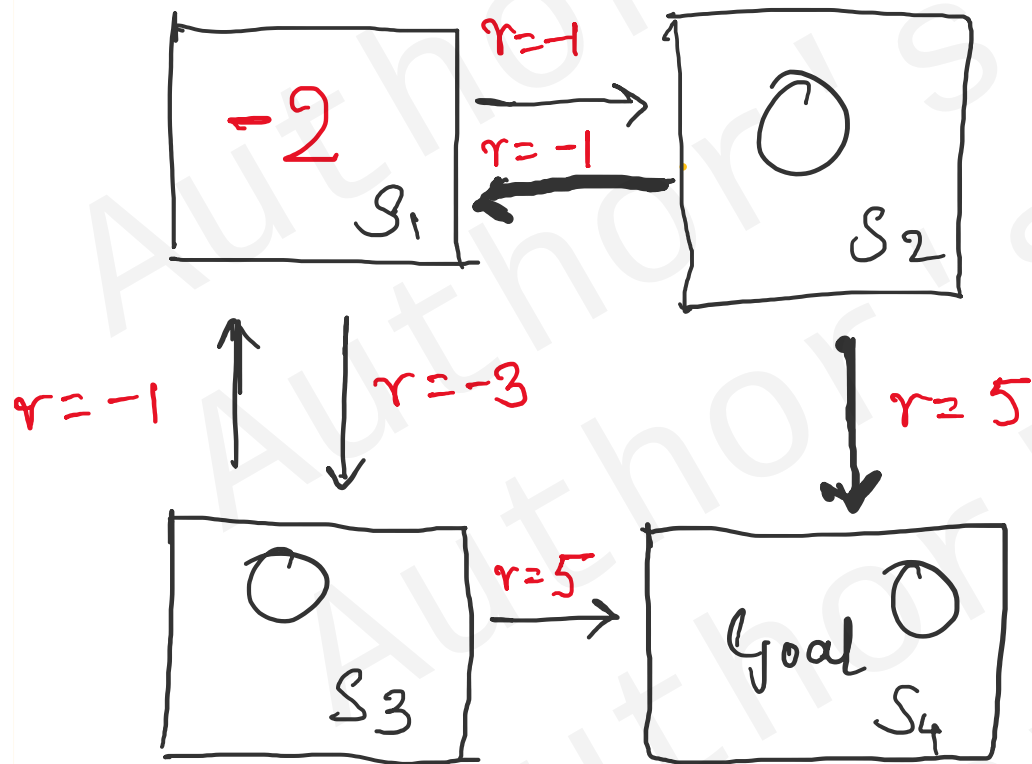


Use  $V_{\pi}(s_1)$  to update value of  $s_2$

$$V_{\pi}(s_2) \leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5)$$

Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

# Grid World Example (Iterative Approach)



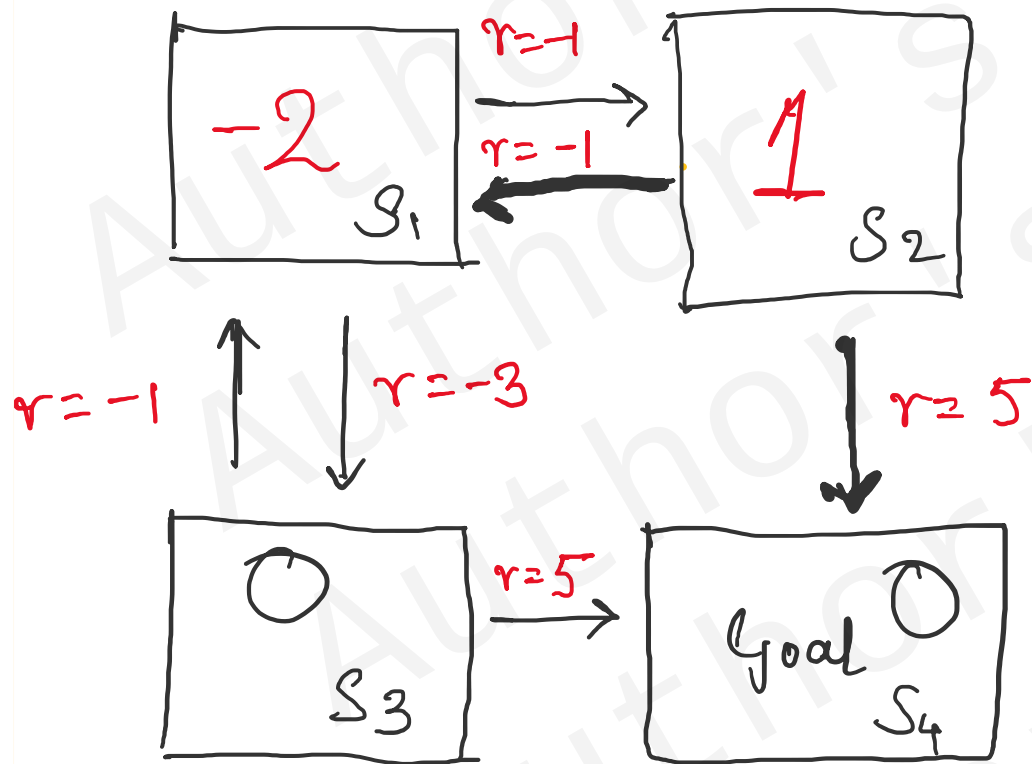
Use  $V_{\pi}(s_1)$  to update value of  $s_2$

$$V_{\pi}(s_2) \leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5)$$

$$V_{\pi}(s_2) \leftarrow 1$$

Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

# Grid World Example (Iterative Approach)



Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

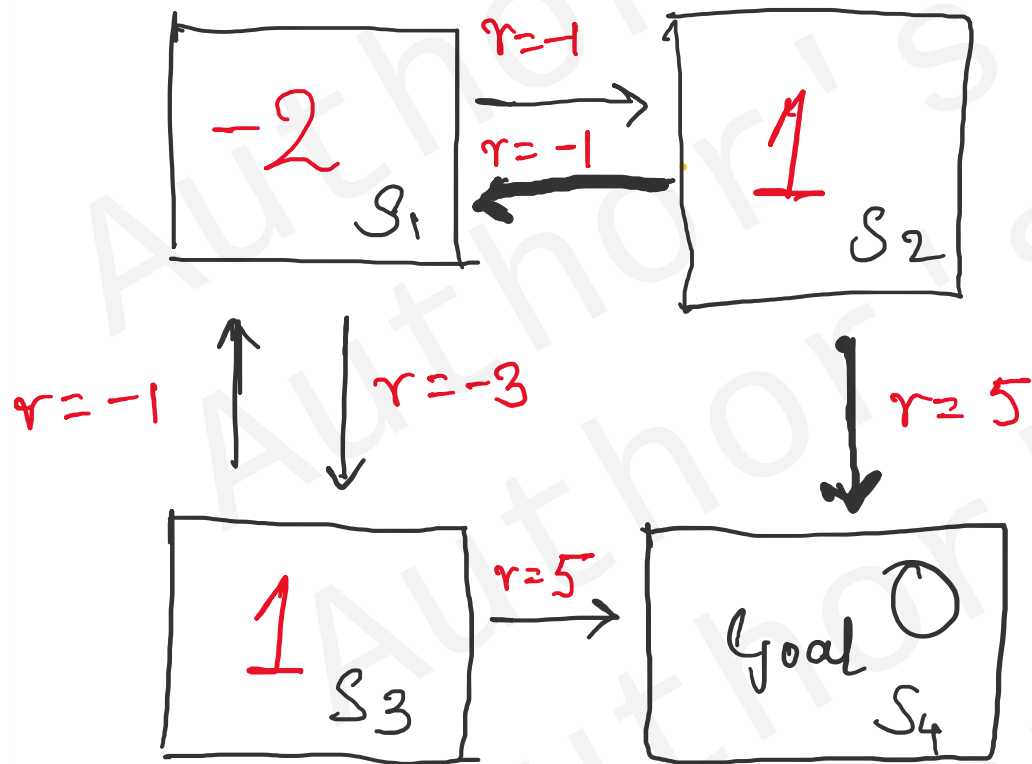
Use  $V_{\pi}(s_1)$  to update value of  $s_2$

$$V_{\pi}(s_2) \leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5)$$

$V_{\pi}(s_2) \leftarrow 1$  record  
keep track and update  $V_{\pi}(s_3)$

$$\begin{aligned} V_{\pi}(s_3) &\leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5) \\ &= \frac{1}{2}(-1 + -2) + \frac{1}{2}(5) = 1 \end{aligned}$$

# Grid World Example (Iterative Approach)



Given  $\pi$ ,  
determine  $V_{\pi}(s) \forall s \in S$

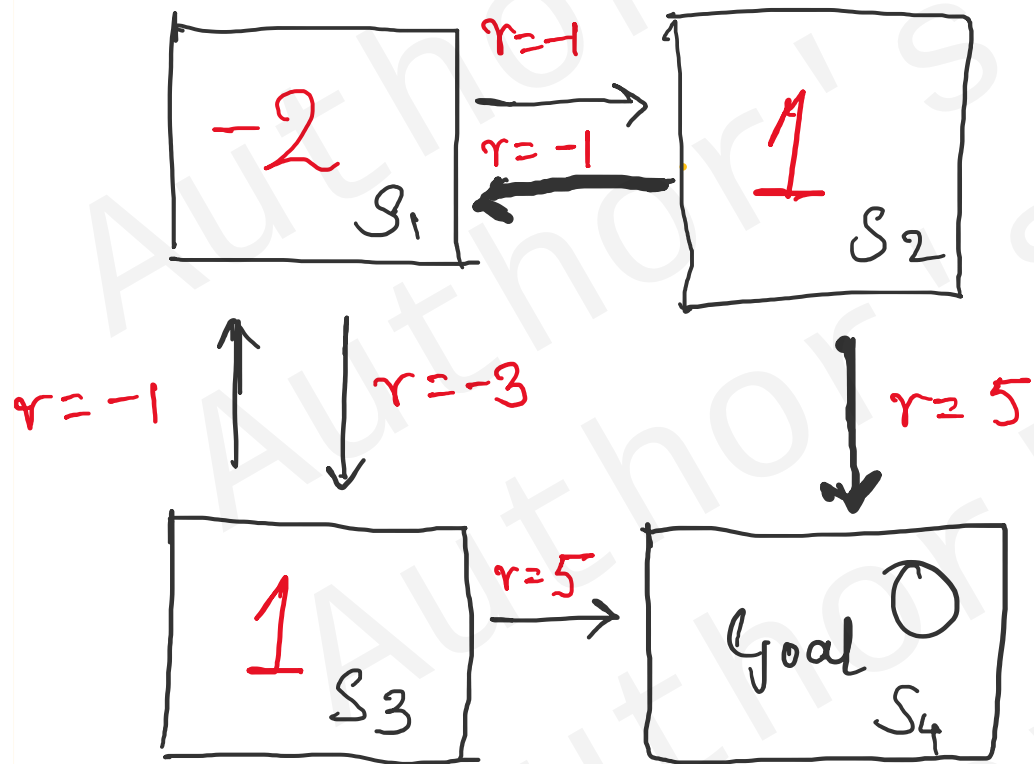
Use  $V_{\pi}(s_1)$  to update value of  $s_2$

$$V_{\pi}(s_2) \leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5)$$

$V_{\pi}(s_2) \leftarrow \underline{1}$  record  
keep track and update  $V_{\pi}(s_3)$

$$\begin{aligned} V_{\pi}(s_3) &\leftarrow \frac{1}{2}(-1 + V_{\pi}(s_1)) + \frac{1}{2}(5) \\ &= \frac{1}{2}(-1 + -2) + \frac{1}{2}(5) = \underline{1} \end{aligned}$$

# Grid World Example (Iterative Approach)



Given  $\pi$ ,  
determine  $V_\pi(s) \forall s \in S$

- Repeat this procedure
- Visit the states over and over
- Converge to True values of states!

Idea behind → Iterative Policy Evaluation

# Iterative Policy Evaluation

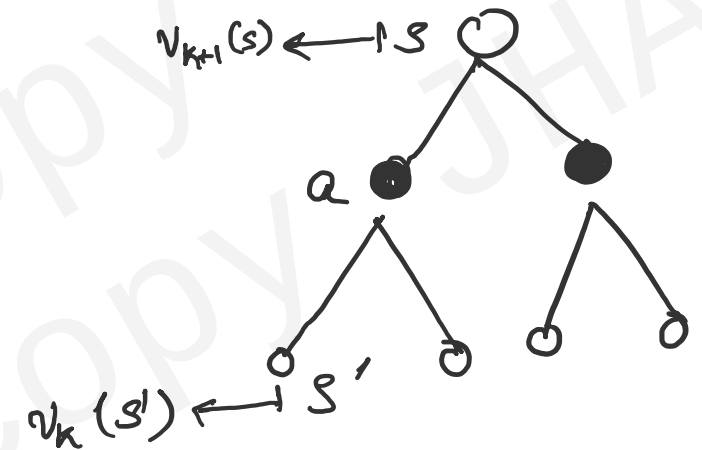
# Iterative Policy Evaluation

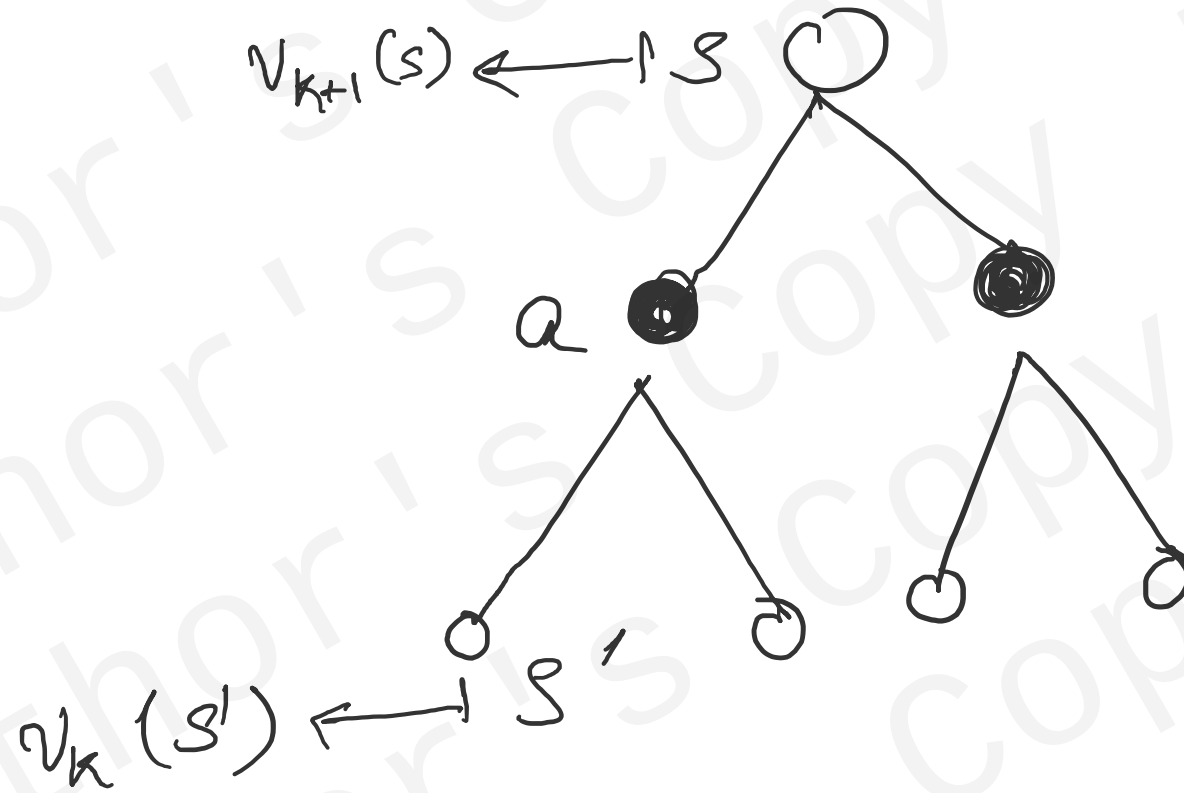
→ Assumes Agent knows MDP, rewards that characterise environment.

Algorithm: to evaluate a given policy  $\pi$

(iterative application of Bellman expectation Backup)

- $v_1 \rightarrow v_2 \rightarrow \dots \rightarrow v_\pi$
- Use synchronous back-ups (storage)
  - At iteration  $k+1$
  - For all states  $s \in S$
  - Update  $v_{k+1}(s)$  from  $v_k(s')$





$$\underline{v}_{k+1}(s) = \sum_{a \in A(s)} \pi(a|s) \sum_{s' \in S, r \in R} p(s', r | s, a) \left( r + \gamma \underline{v}_k(s') \right)$$

for all  $s \in S$ .

# Iterative Policy Evaluation

Iterative algorithm adapts the Bellman Eq to be used as update rule:

$$V(s) \leftarrow \sum_{a \in A} \pi(a/s) \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

↑  
Value function

↑  
Value function.

# Iterative Policy Evaluation

Iterative algorithm

Input: MDP, policy  $\pi$

Output: V-function ( $V_\pi$ )

$V(s) \leftarrow 0$  for all  $s \in S$

repeat:  
for  $s \in S$

$$V(s) \leftarrow \sum_{a \in A} \pi(a/s) \sum_{s' \in S, r \in R} p(s', r/s, a) (r + \gamma V(s'))$$

# Iterative Policy Evaluation

Iterative algorithm

Input: MDP, policy  $\pi$

Output: V-function ( $V_\pi$ )

$V(s) \leftarrow 0$  for all  $s \in S$  | begin with initial guess

repeat:

for  $s \in S$  | loop over all states

| update the values

$$V(s) \leftarrow \sum_{a \in A} \pi(a/s) \sum_{s' \in S, r \in R} p(s', r/s, a) (r + \gamma V(s'))$$

If value fn  $\rightarrow$  contraction mapping, state  $\rightarrow$  compact set  
guaranteed to converge to  $V_\pi$ .

# Iterative Policy Evaluation

- Guaranteed to converge to  $V_\pi$  in infinite sense.
- But how to assess if convergence is reached in approximate sense?
- Introduce a  $\theta$

# Iterative Policy Evaluation

Input: MDP, policy  $\pi$ , small positive  $\theta$

Output: V-function ( $V_\pi$ )

$V(s) \leftarrow 0$  for all  $s \in S$

| begin with initial guess

repeat untill  $\Delta < \theta$ :

$\Delta \leftarrow 0$

for  $s \in S$  | loop over all states

$v \leftarrow V(s)$

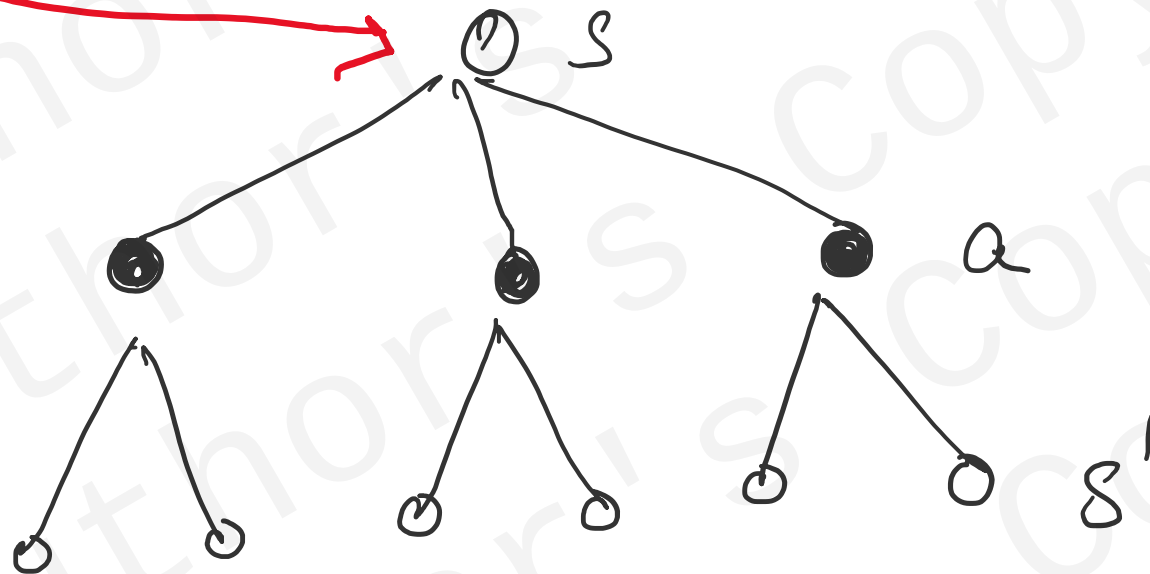
$$V(s) \leftarrow \sum_{a \in A} \pi(a/s) \sum_{s' \in S, r \in R} p(s', r/s, a) (r + \gamma V(s'))$$

$\Delta \leftarrow \max(\Delta, |v - V(s)|)$

| how much values of each state changed  
if not changed, stop.

return  $V$

$$V(s) \leftarrow \sum_{a \in A} \pi(a/s) \sum_{s' \in S, r \in R} p(s', r/s, a) (r + \gamma V(s'))$$



# Iterative Policy Evaluation

Policy  $\pi$



Value function  $V_\pi$

We have  $V_\pi$



But, optimal  $V$ -function  
still not found!!

# Policy Improvement

# How to improve a Policy

## Policy Evaluation

policy  $\pi$



value function  
 $v_\pi$

## Policy Improvement

value function  $v_\pi$



policy  $\pi'$  (with  $\pi' \geq \pi$ )

policy  $\pi$



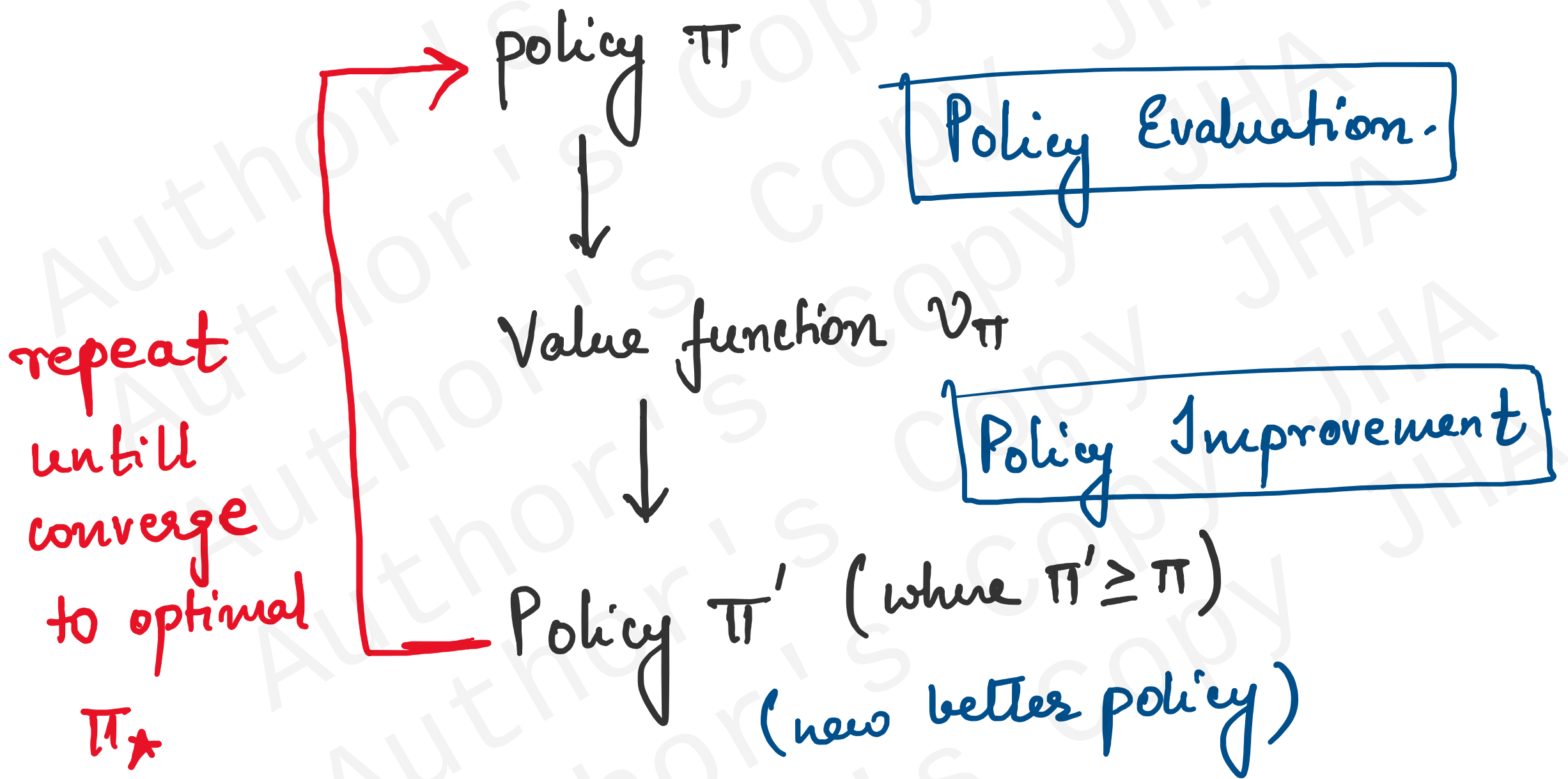
Value function  $V_\pi$

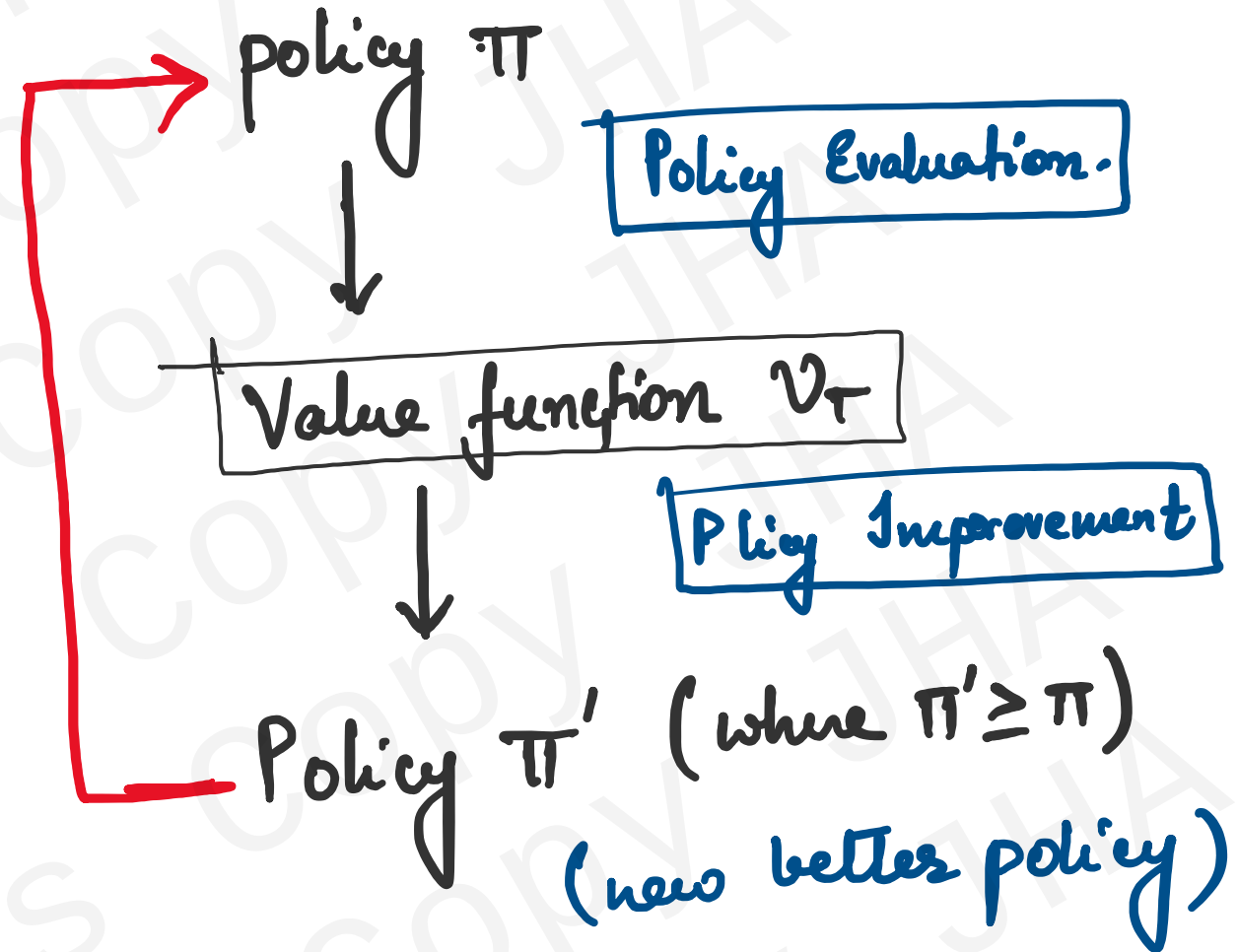
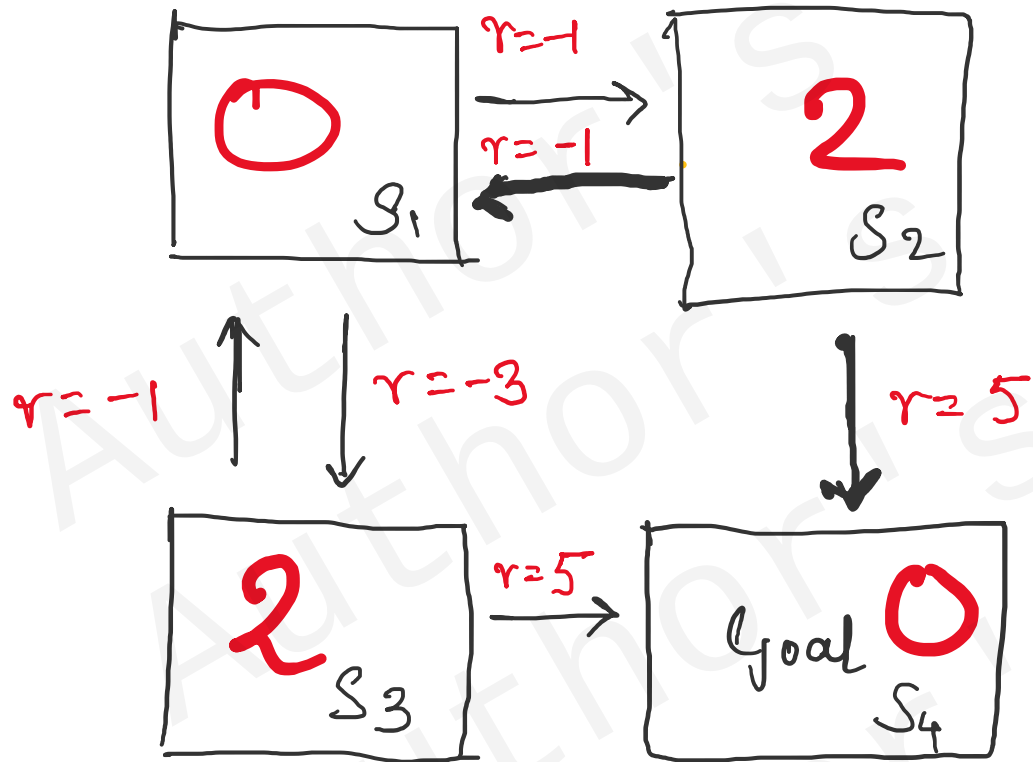


Policy  $\pi'$  (where  $\pi' \geq \pi$ )  
(new better policy)

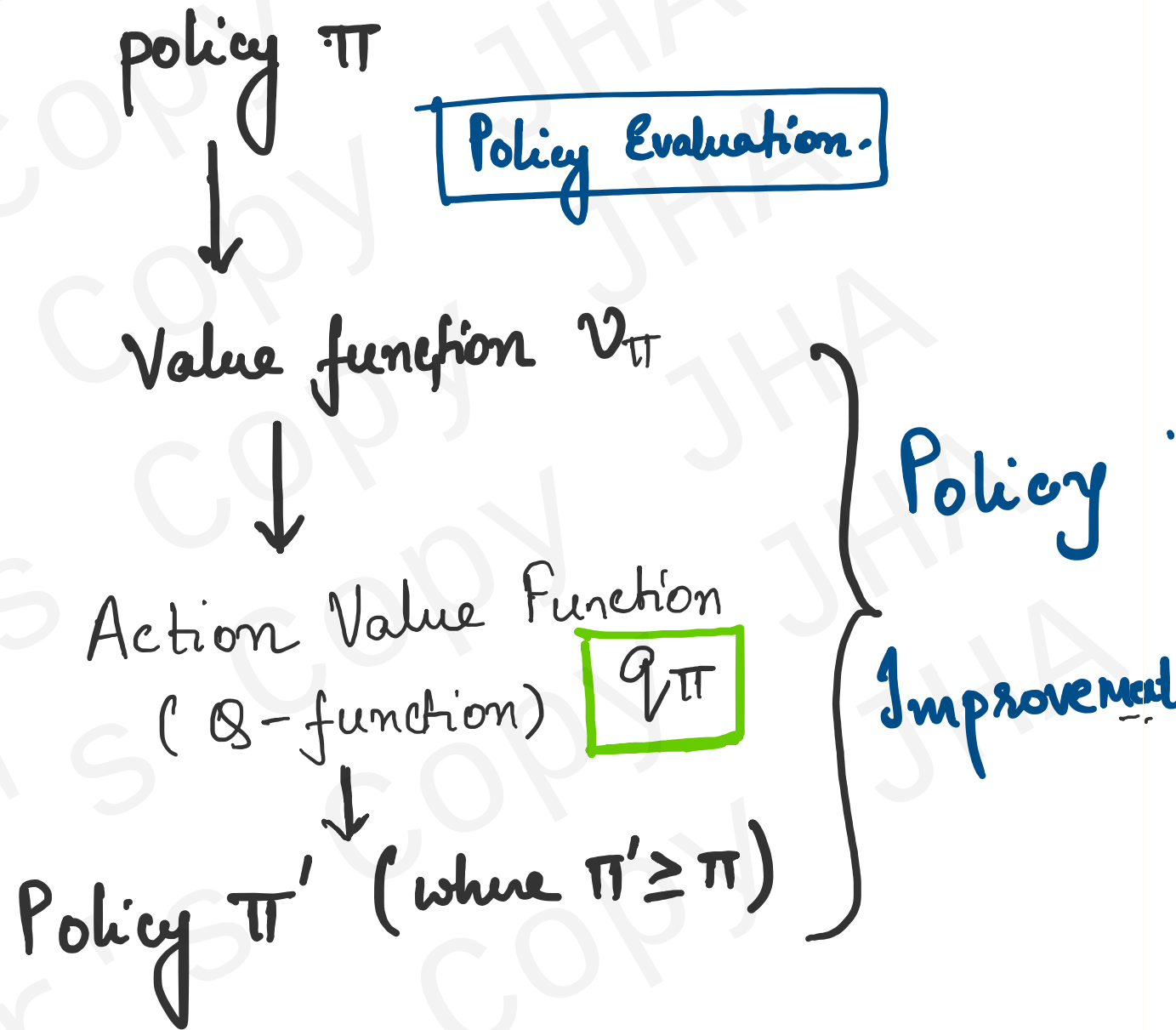
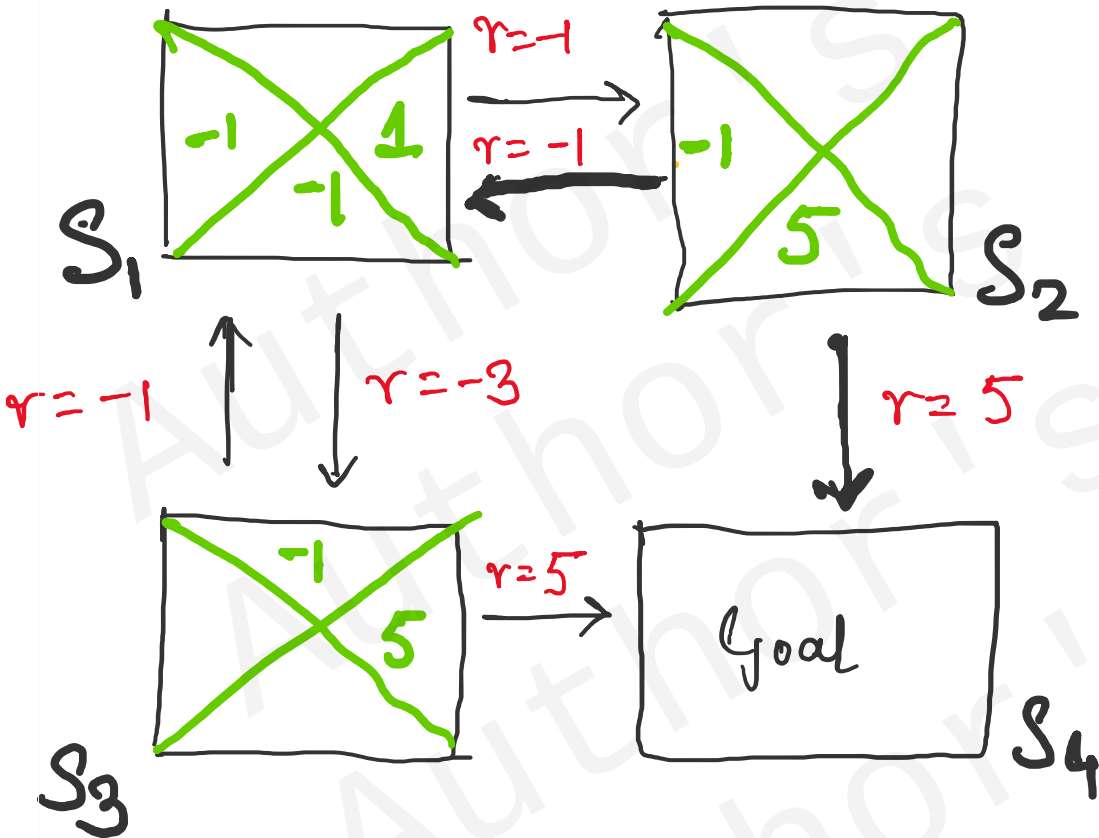
Policy Evaluation

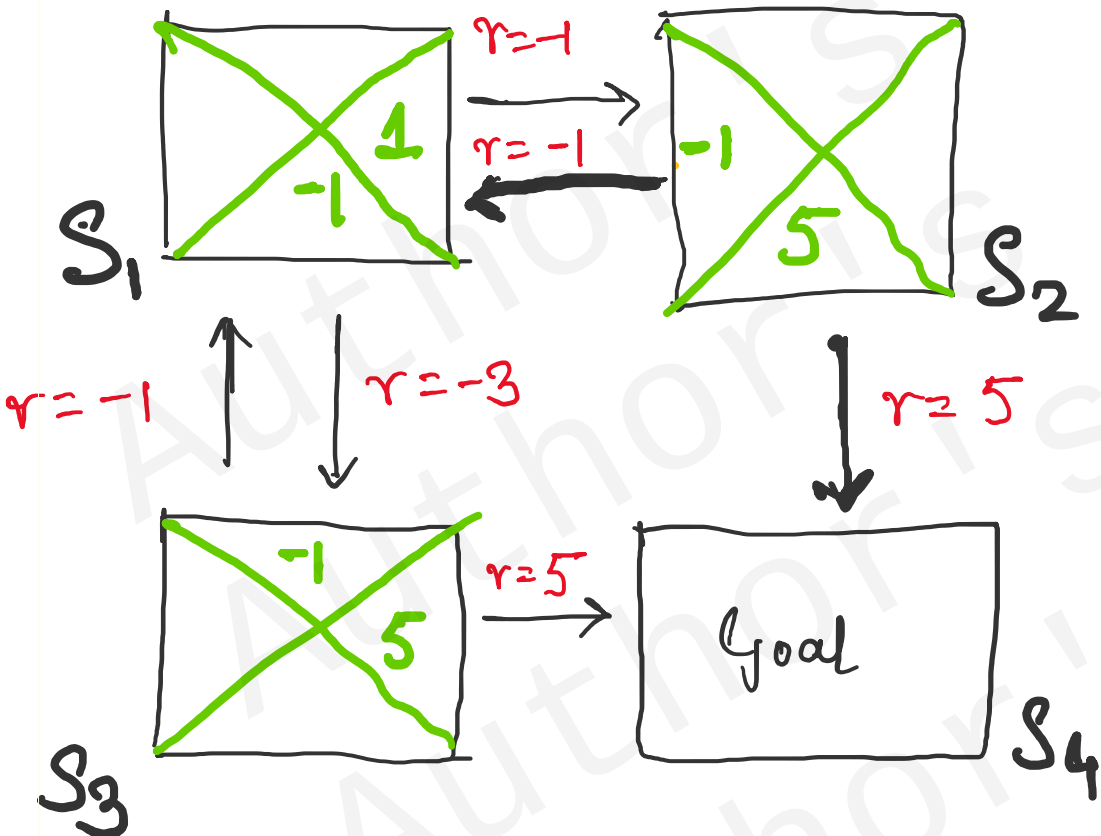
Policy Improvement



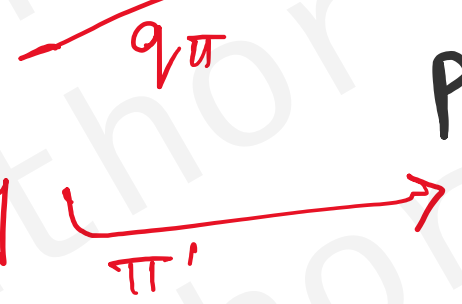


Now, how to improve policy.  
 choose policy  $\pi' = \text{greedy}(\pi)$  → gives better value for all  $V_\pi$ .





Focus: How to use  
to choose better policy



policy  $\pi$

Policy Evaluation.

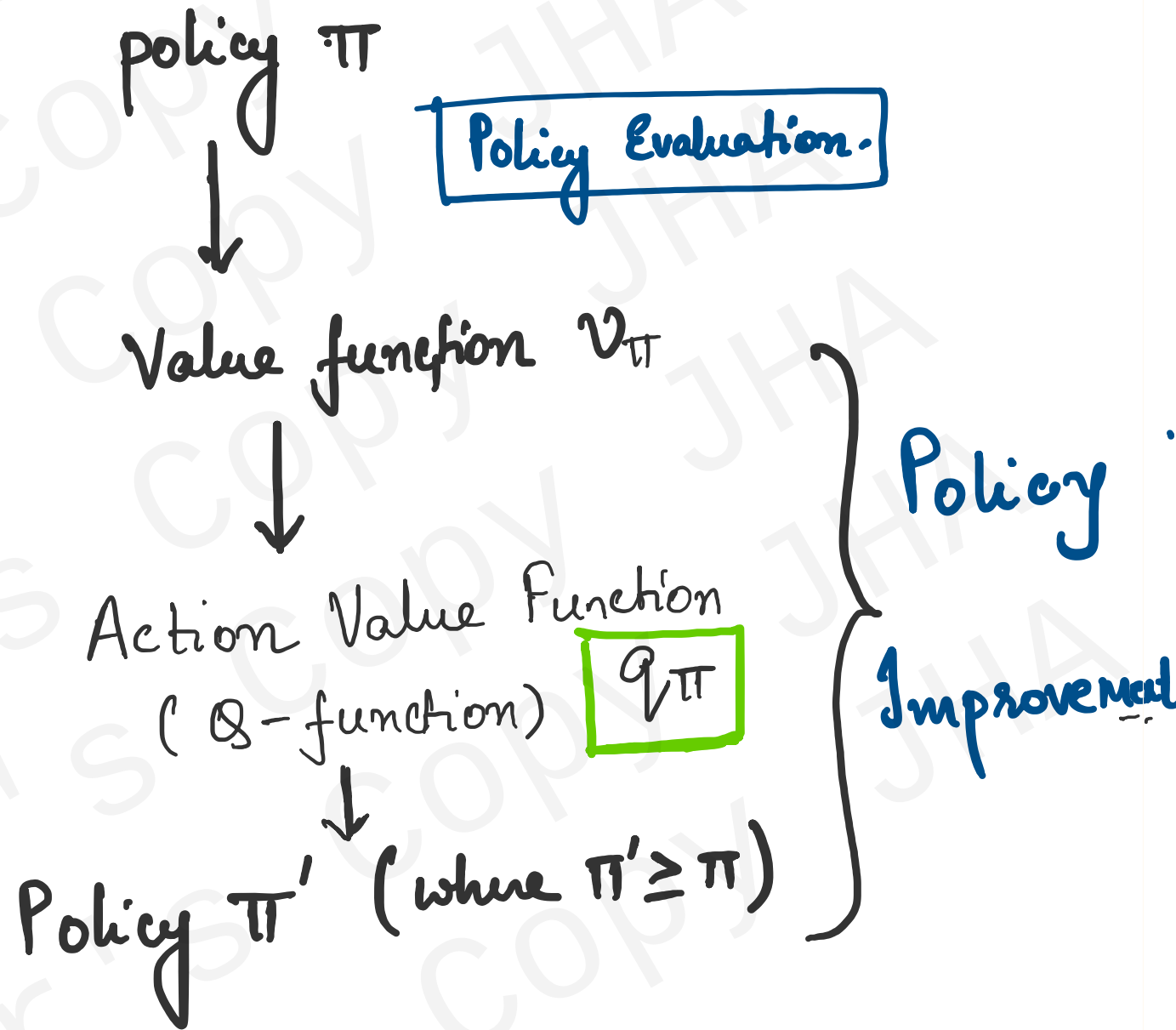
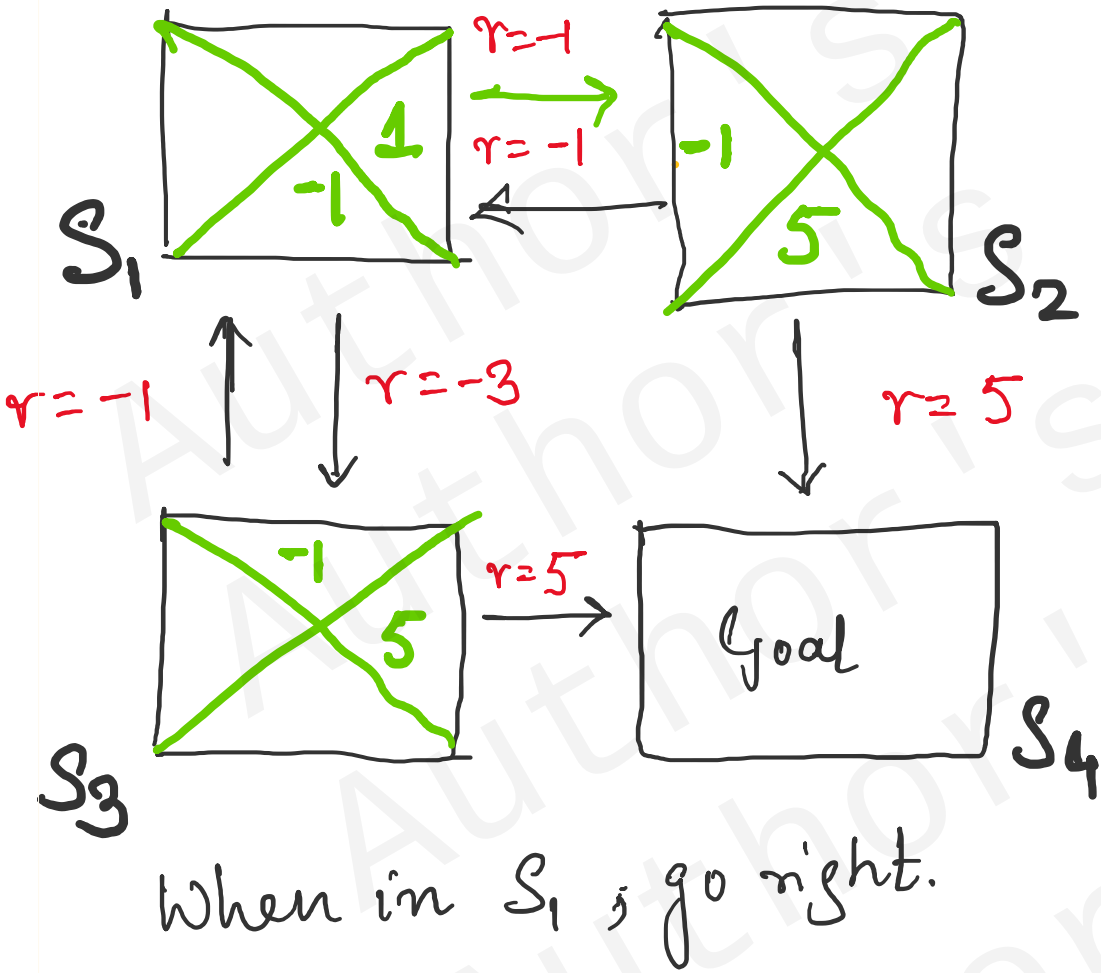
Value function  $V_\pi$

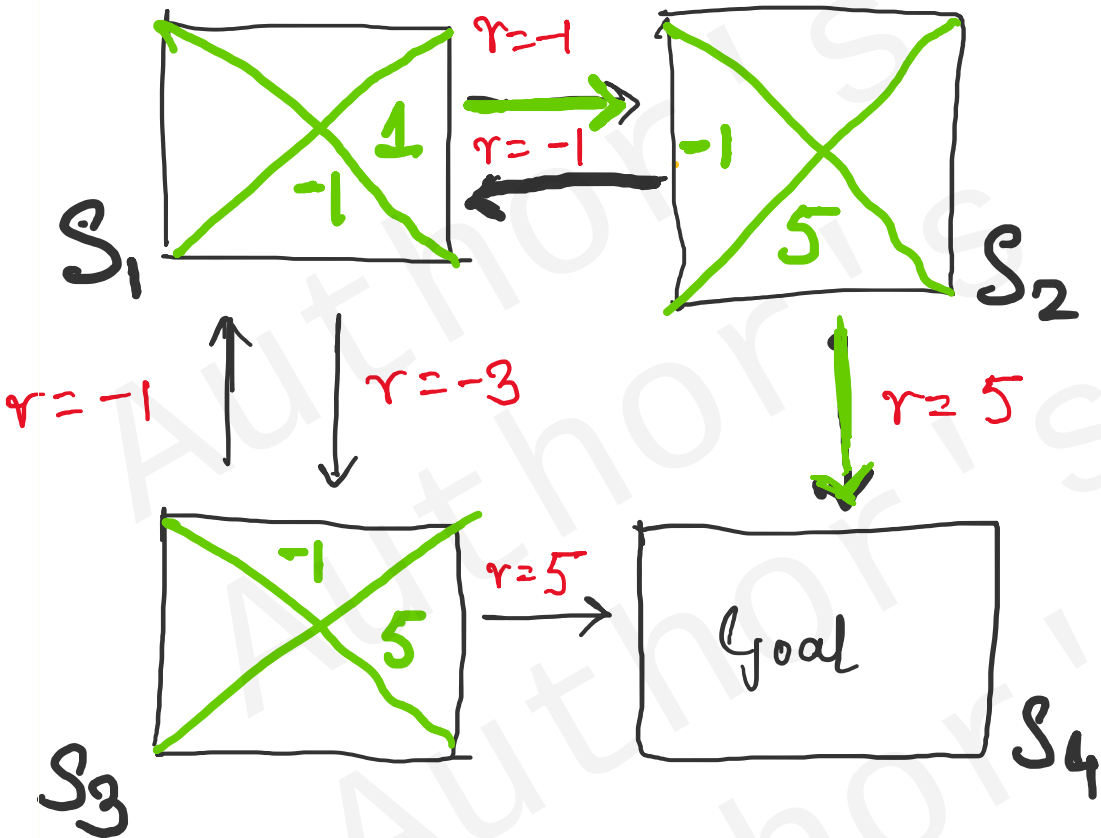
Action Value Function  
(Q-function)  $Q_\pi$

Policy  $\pi'$  (where  $\pi' \geq \pi$ )

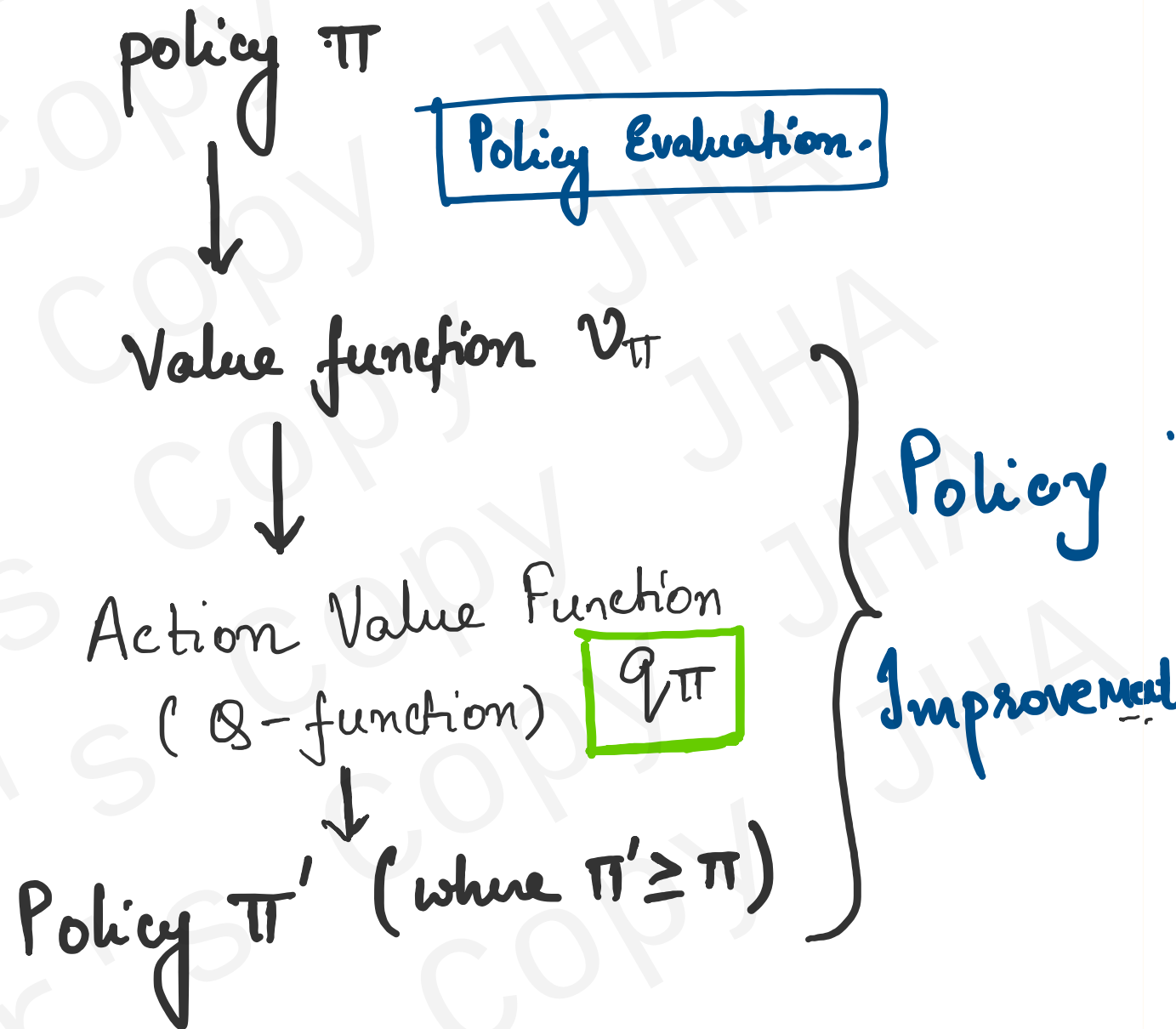
Policy

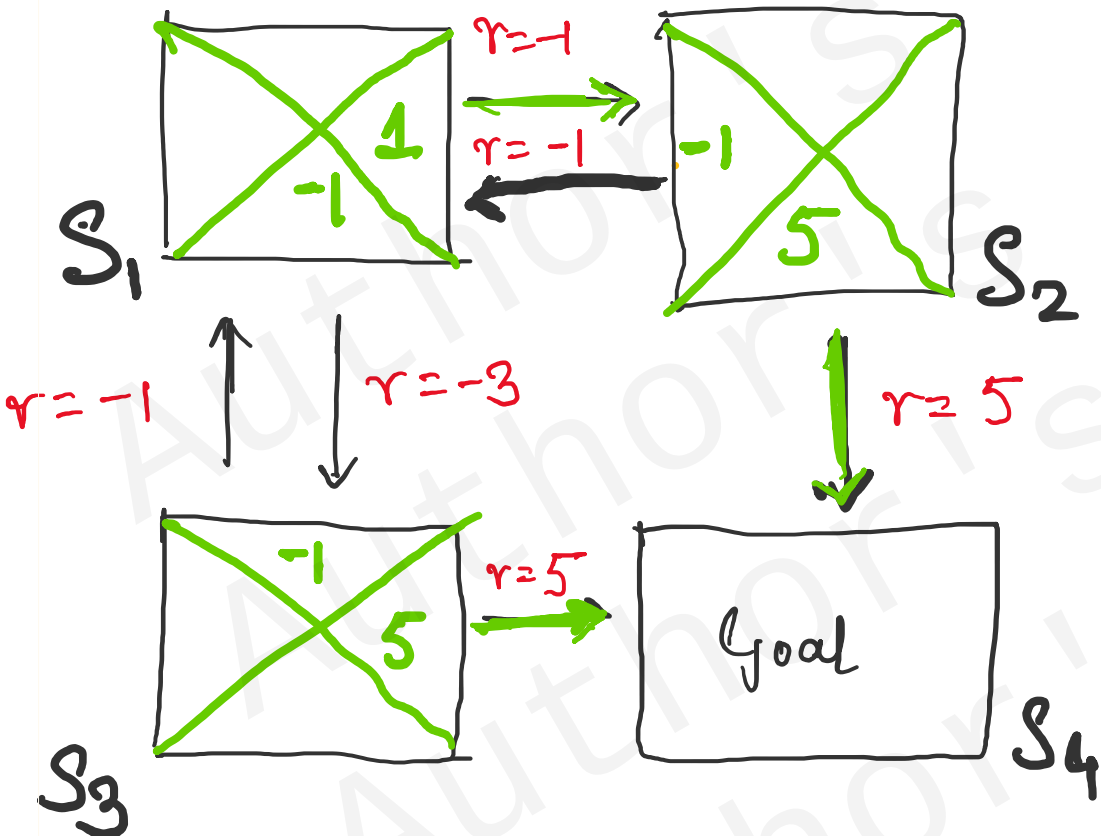
Improvement





When in  $S_2$ , go down.





When in  $S_3$ , go right.

Optimal policy in Green. ←

policy  $\pi$

**Policy Evaluation.**

Value function  $V_\pi$

Action Value Function  
(Q-function)  **$Q_\pi$**

Policy  $\pi'$  (where  $\pi' \geq \pi$ )

(if and only if,  $V_{\pi'}(s) \geq V_\pi(s)$  for all  $s \in S$ )

Policy Improvement

# Policy Improvement.

- Consider a policy  $a = \pi(s)$
- Improve the policy by greedy approach.

$$\pi'(s) = \operatorname{argmax}_{a \in A} q_{\pi}(s, a)$$

- This improves value from any state  $(s)$  over one step.

$$q_{\pi}(s, \pi'(s)) = \max_{a \in A} q_{\pi}(s, a) \geq q_{\pi}(s, \pi(s)) = v_{\pi}(s)$$

# Policy Improvement.

Value function  $V_{\pi'}(s) \geq V_{\pi}(s)$   
is improved.

$$\begin{aligned} V_{\pi}(s) &\leq q_{\pi}(s, \pi'(s)) = \mathbb{E}_{\pi'}(R_{t+1} + \gamma V_{\pi}(s_{t+1}) \mid S_t = s) \\ &\leq \mathbb{E}_{\pi'}(R_{t+1} + \gamma q_{\pi}(s_{t+1}, \pi'(s_{t+1}))) \mid S_t = s \\ &\leq \mathbb{E}_{\pi'}(R_{t+1} + \gamma R_{t+2} + \gamma^2 q_{\pi}(s_{t+2}, \pi'(s_{t+2}))) \mid S_t = s \\ &\leq \mathbb{E}_{\pi'}(R_{t+1} + \gamma R_{t+2} + \dots \mid S_t = s) = V_{\pi'}(s) \end{aligned}$$

Thus,  $V_{\pi}(s) \leq V_{\pi'}(s)$

# Policy Improvement

→ If improvements stop,

$$q_{\pi}(s, \pi'(s)) = \max_{a \in A} q_{\pi}(s, a) = q_{\pi}(s, \pi(s)) = v_{\pi}(s)$$

→ Then, Bellman Optimality is stopped,

$$v_{\pi}(s) = \max_{a \in A} q_{\pi}(s, a)$$

→ Therefore,  $v_{\pi}(s) = v_{*}(s)$  for all  $s \in S$

π is optimal policy

Policy iteration algorithm

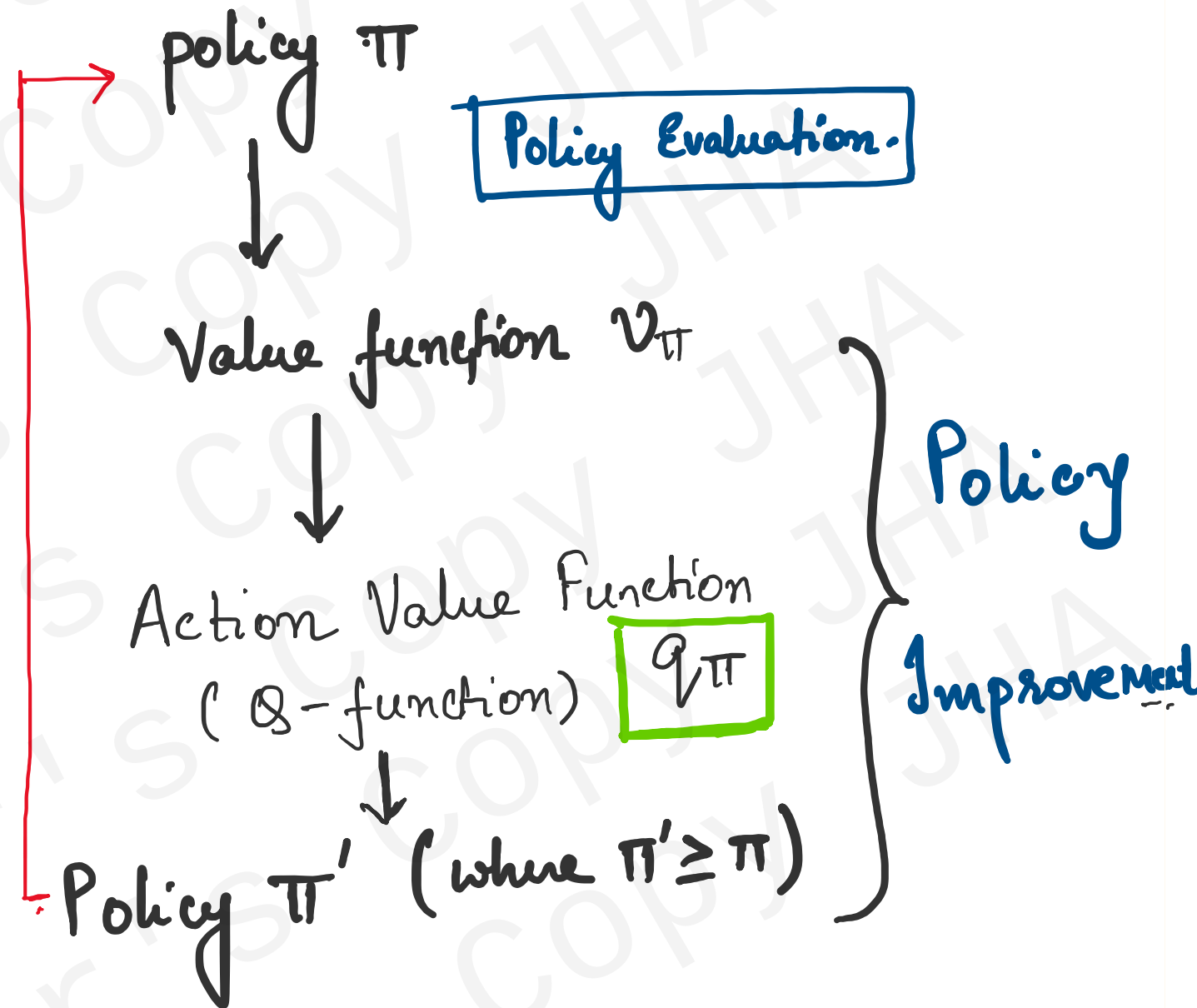
=

Policy evaluation

+

Policy improvement

# Policy Iteration



# Policy Iteration

=

## Policy Evaluation

+

## Policy Improvement

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

for  $s \in S$

for  $a \in A(s)$

$$Q(s, a) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)} Q(s, a)$$

# Policy Iteration Algorithm

**Input:** MDP, small positive number  $\epsilon$

**Output:** policy  $\pi = \pi^*$

$$\pi(a/s) = \frac{1}{|A(s)|} \text{ for all } s \in S \text{ and } a \in A(s)$$

begin with equiprobable random policy

repeat:

$V \leftarrow$  Policy-Evaluation (MDP,  $\pi$ ,  $\epsilon$ )

obtain V-fun wrt  $\pi$

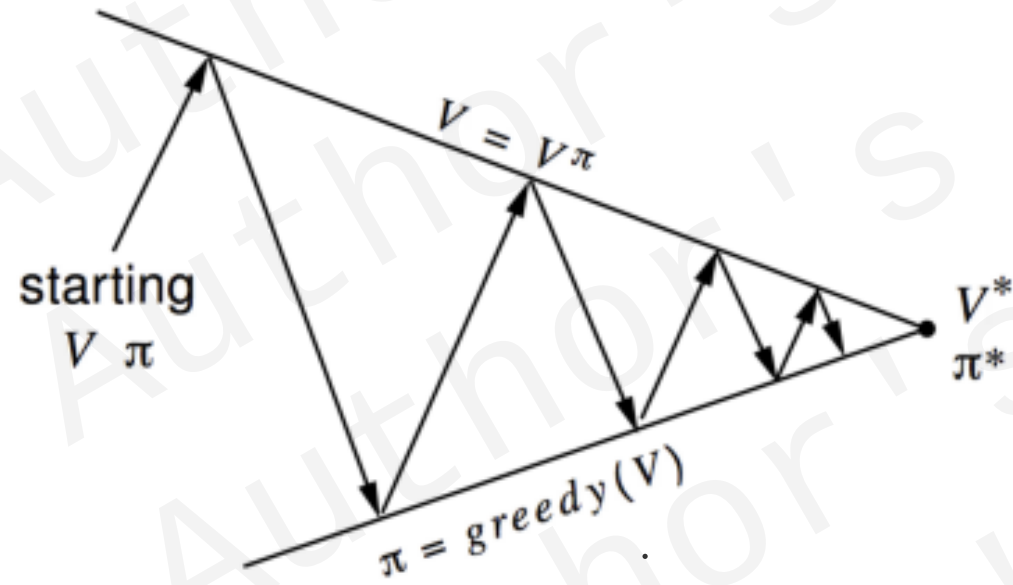
$\pi' \leftarrow$  Policy-Improvement (MDP,  $V$ )

obtain better or equal policy from V-function

if  $\pi = \pi'$  then Break

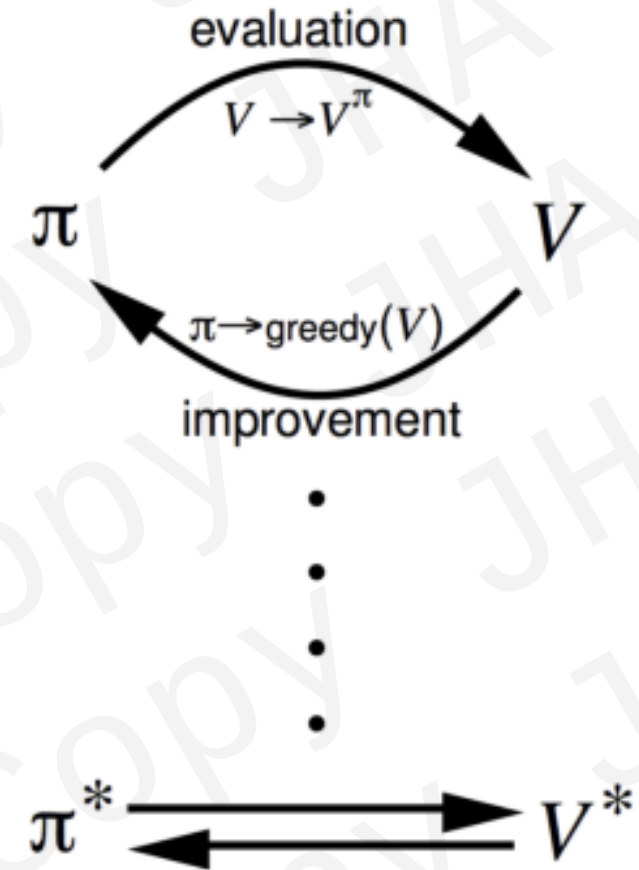
return  $\pi$

# Policy Iteration

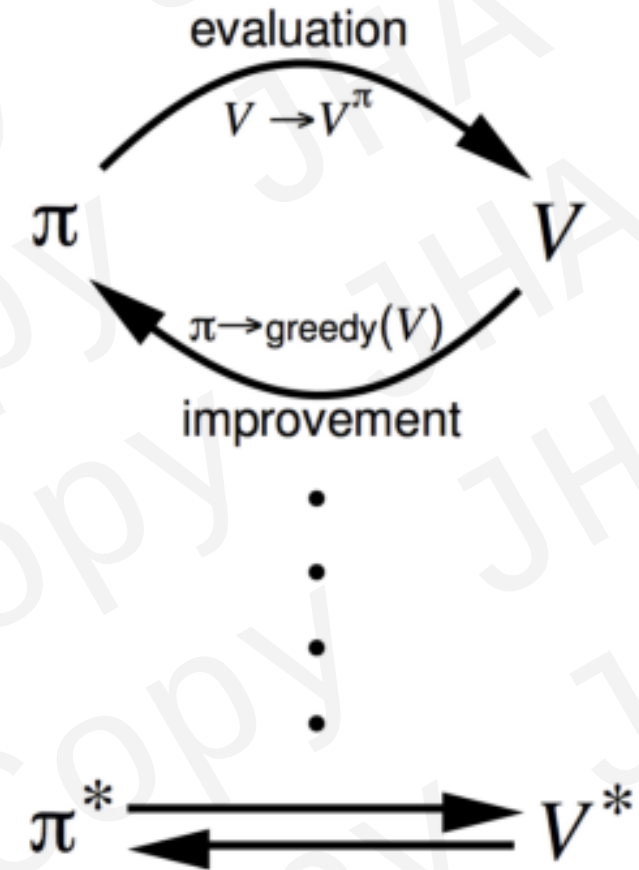
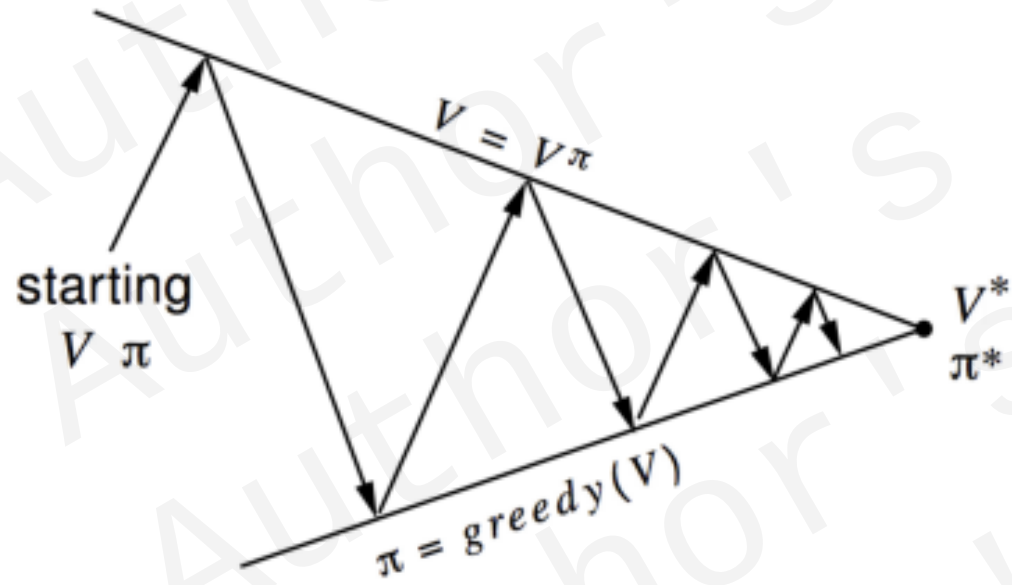


Policy Evaluation : Estimate  $V^\pi$

Policy Improvement : Greedy ( $\pi' \geq \pi$ )



# Generalised Policy Iteration



# Principle of Optimality

Optimal policy = optimal first action  $A^*$  + optimal policy from successor state  $s'$ .

Theorem: A policy  $\pi(a|s)$  achieves optimal value from state  $s$ ,

$$V_\pi(s) = V_*(s) \text{ if and only if}$$

– any state  $s'$  is reachable from  $s$

–  $\pi$  achieves optimal value from  $s'$

$$\text{or } V_\pi(s') = V_*(s')$$

## Value Iteration algorithm

# ~~Policy Iteration~~ Value Iteration

Policy Evaluation  
(one time)

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Policy Improvement.

for  $s \in S$

for  $a \in A(s)$

$$Q(s, a) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)} \underline{Q(s, a)}$$

# ~~Policy Iteration~~ Value Iteration

Policy Evaluation  
(one time) for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Policy Improvement  
for  $s \in S$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

# ~~Policy Iteration~~ Value Iteration

Policy Evaluation for  $s \in S$   
(one time)

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Policy Improvement:

for  $s \in S$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

Policy Evaluation:

$$\pi \leftarrow \pi'$$

(single sweep)

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Redundancy Here

# Policy Iteration

Policy Evaluation for  $s \in S$   
(one time)

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Policy Improvement:

for  $s \in S$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

$$\pi \leftarrow \pi'$$

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Policy Evaluation:

(single sweep)

Redundancy Here

for  $s \in S$

$$\pi'(s) \leftarrow \arg \max_{a \in A(s)}$$

$$\sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

$$\pi \leftarrow \pi'$$

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

Identical Expression except the actions.

$f(a)$  ①

for  $s \in S$

$$\pi(s) \leftarrow \arg \max_{a \in A(s)}$$

$$\sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

②  $\pi(s)$  is  $\arg \max$  over 'a'

for  $s \in S$

$$V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s'))$$

$f(\pi(s))$  ③

↑ evaluated at  $\pi(s)$

Combining ①, ② and ③  
(set  $V(s)$  to max of the function)

$$\begin{aligned} & \text{for } s \in S \\ & \pi(s) \leftarrow \arg \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s')) \\ & \text{for } s \in S \\ & V(s) \leftarrow \sum_{s' \in S, r \in R} p(s', r | s, \pi(s)) (r + \gamma V(s')) \\ & V(s) \leftarrow \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | s, a) (r + \gamma V(s')) \end{aligned}$$

# Value Iteration

Input: MDP, small positive  $\theta$

Output: policy  $\pi^*$

$V(s) \leftarrow 0$  for all  $s \in S$   
Repeat until  $\Delta < \theta$ .

start with some initial guess

$$\Delta \leftarrow 0$$

for  $s \in S$ :

$$v \leftarrow V(s)$$

$$V(s) \leftarrow \max_{a \in A(s)} \sum_{s', r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

$$\Delta \leftarrow \max(\Delta, |v - V(s)|)$$

loop over all states

update  
V-fun values

stop when not  
updating

for  $s \in S$

$$\pi(s) = \arg \max_{a \in A(s)} \sum_{s', r \in R} p(s', r | s, a) (r + \gamma V(s'))$$

Obtain policy  
with respect to final value  
function.

# Real Time Dynamic Programming.

- 1) Use states relevant to the agent
- 2) Use agent's experience to guide the selection of states
- 3) After each time step:  $S_t, A_t, R_{t+1}$
- 4) Back up the state  $S_t$

$$v(S_t) \leftarrow \max_{a \in A(s)} \sum_{s' \in S, r \in R} p(s', r | S_t, a) (r + \gamma v(s'))$$

Here  $s'$  is the state at time  $t+1$  or  $s' = S_{t+1}$

# Back Ups.

- 1) DP uses full back ups
- 2) Every successor state is considered
- 3) Knowledge of MDP and reward transition is used
- 4) For large problems, large state space  
'curse of dimensionality'  
number of states  $n = |S|$   
↓  
grows exponentially with  
number of variables.